



Analisis Prediksi Customer *Churn* Nasabah Bank dengan Algoritma Random Forest

Anindita Puspa Ayu Prayogi ^{a,1,*}, Sucipto ^{b,2}, Arie Nugroho ^{b,3}

^a Universitas Nusantara PGRI Kediri, Jl. Ahmad Dahlan No.76 Mojoroto, Kediri, Indonesia

^b Universitas Nusantara PGRI Kediri, Jl. Ahmad Dahlan No.76 Mojoroto, Kediri, Indonesia

¹ aninditapuspa3@gmail.com ^{*}; ² sucipto@unpkediri.ac.id; ³ arienugroho@unpkediri.ac.id

^{*} Corresponding Author

ARTICLE INFO

Article History

Submission : 17-10-2025

Revision : 29-10-2025

Accepted : 08-03-2026

Kata Kunci:

churn, Random Forest, *feature selection*, perbankan, klasifikasi nasabah

ABSTRAK

Tingkat *churn* nasabah yang tinggi merupakan tantangan serius bagi sektor perbankan karena berdampak pada pendapatan dan loyalitas jangka panjang. Penelitian ini bertujuan untuk meningkatkan akurasi klasifikasi *churn* nasabah bank menggunakan algoritma Random Forest serta mengidentifikasi faktor-faktor utama yang memengaruhi *churn*. Berbeda dengan penelitian terdahulu yang umumnya menggunakan konfigurasi standar, penelitian ini menerapkan Random Forest dengan penyesuaian parameter serta analisis *feature importance* untuk mengevaluasi dampak *feature selection* terhadap performa model. Dataset publik dari Kaggle yang terdiri dari 10.000 data nasabah digunakan dalam penelitian ini. Dua model Random Forest dibangun, yaitu model dengan seluruh fitur dan model dengan fitur terseleksi. Hasil evaluasi menunjukkan bahwa model tanpa *feature selection* menghasilkan akurasi tertinggi sebesar 86,7%, sementara penerapan *feature selection* tidak meningkatkan performa model. Faktor usia, jumlah produk, dan saldo rekening merupakan variabel paling berpengaruh terhadap *churn*. Hasil penelitian ini menegaskan keunggulan Random Forest dalam prediksi *churn* nasabah bank melalui pendekatan evaluatif terhadap *feature selection* serta memberikan kontribusi berupa analisis komparatif performa model yang dapat menjadi dasar pengambilan keputusan strategis dalam manajemen retensi nasabah.

This is an open access article under license [CC-BY-SA](#).



1. Pendahuluan

Dalam dunia bisnis modern, mempertahankan pelanggan menjadi kunci utama bagi keberlangsungan dan pertumbuhan perusahaan [1]. Loyalitas pelanggan tidak hanya memberikan keuntungan dari transaksi berulang, tetapi juga memperkuat reputasi melalui rekomendasi positif. Untuk bersaing dalam dunia bisnis yang semakin kompetitif, perusahaan di berbagai sektor berupaya menciptakan pengalaman pelanggan yang memuaskan melalui strategi kepuasan pelanggan yang efektif [2]. Oleh karena itu, perusahaan perlu menerapkan

strategi adaptasi, distribusi, promosi, dan riset menyeluruh untuk memahami perilaku dan kebutuhan pelanggan [2].

Dalam sektor perbankan, menjaga loyalitas nasabah menjadi semakin kompleks karena tingginya tingkat persaingan dan kemajuan teknologi yang memungkinkan nasabah dengan mudah berpindah ke layanan lain, yang disebut sebagai *churn* [3]. *Customer churn* atau *customer attrition* merupakan fenomena ketika nasabah mengakhiri hubungan dengan suatu organisasi, yang dalam konteks perbankan terjadi saat nasabah menutup rekening atau menghentikan penggunaan layanan, sehingga pemahaman dan pengelolaan *churn* menjadi krusial untuk menjaga stabilitas keuangan dan reputasi bank [4]. Oleh sebab itu, bank perlu mengidentifikasi dan memprediksi faktor-faktor yang mendorong *churn* nasabah, karena potensi kerugian yang ditimbulkan cukup besar, sehingga dibutuhkan pemanfaatan analisis data untuk memahami perilaku nasabah dan merancang strategi retensi yang lebih efektif [5].

Di era digital saat ini, *Artificial Intelligence* (AI) berperan penting dalam analisis data, dengan data mining dan *machine learning* menjadi teknik klasifikasi yang populer di berbagai bidang [6]. Aktivitas manusia menghasilkan beragam data dalam berbagai ukuran yang tidak akan bernilai tanpa pengolahan yang tepat, sehingga diperlukan penerapan teknik data mining dan *machine learning* untuk mengubah dataset historis menjadi informasi dan pengetahuan yang bermanfaat bagi aktivitas manusia di masa depan [7]. Untuk mengatasi hal tersebut, berbagai pendekatan berbasis data telah dikembangkan, termasuk penerapan teknik data mining dan *machine learning*. Data mining merupakan salah satu bidang dalam *Artificial Intelligence* yang memungkinkan perusahaan menganalisis pola dan tren kompleks dalam perilaku nasabah [8]. Machine learning berperan penting dalam retensi pelanggan dengan memantau perubahan perilaku nasabah dan memprediksi customer churn, sehingga perusahaan dapat mengidentifikasi pelanggan berisiko tinggi dan mengambil tindakan proaktif sebelum pelanggan benar-benar meninggalkan layanan [4].

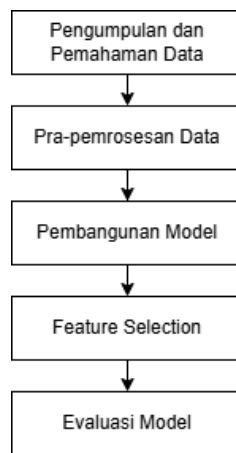
Algoritma klasifikasi seperti Naive Bayes dan ID3 telah digunakan dalam penelitian sebelumnya [3]. Klasifikasi adalah proses mengenali pola untuk membedakan data ke dalam kelompok tertentu, sehingga setiap objek dapat dimasukkan ke kategori yang sesuai berdasarkan kriteria yang ditetapkan [9]. Meskipun hasil dari algoritma-algoritma ini menunjukkan akurasi yang cukup baik, masih ada potensi untuk meningkatkan akurasi model. Penelitian ini menggunakan algoritma Random Forest, yang dikenal unggul dalam mengolah data kompleks serta menyediakan analisis *feature importance* untuk memahami variabel-variabel yang memengaruhi keputusan *churn*.

Random forest merupakan salah satu metode pembelajaran mesin yang paling banyak digunakan untuk masalah prediksi dan klasifikasi. Untuk membuat model dengan *Random Forest*, *dataset* dipecah menjadi beberapa bagian untuk membangun model *decision tree*, kemudian *data testing* diuji satu per satu pada setiap *decision tree*. Setiap *decision tree* akan memberikan hasil atau *vote* berdasarkan model masing-masing [10]. Nilai tertinggi dari *vote* pada setiap *decision tree* akan dijadikan hasil akhir proses klasifikasi. Selain membangun model klasifikasi, penelitian ini juga memanfaatkan proses *feature selection* untuk mengevaluasi dampaknya terhadap performa model.

Tujuan utama dari penelitian ini adalah untuk meningkatkan akurasi model klasifikasi *churn* nasabah bank menggunakan algoritma Random Forest dan mengidentifikasi fitur-fitur yang paling berpengaruh terhadap keputusan nasabah dalam berhenti menggunakan layanan bank. Dalam konteks ini, pemanfaatan algoritma yang mampu menangani data kompleks dan memberikan interpretasi yang jelas terhadap faktor-faktor *churn* menjadi sangat penting. Oleh karena itu, penggunaan Random Forest tidak hanya ditujukan untuk meningkatkan akurasi, tetapi juga untuk memahami perilaku nasabah secara lebih mendalam melalui analisis kontribusi fitur.

2. Metode Penelitian

Penelitian ini bertujuan untuk memprediksi atau mengklasifikasikan *customer churn*, yaitu kondisi ketika nasabah memutuskan untuk berhenti menggunakan layanan bank. Fenomena ini dipengaruhi oleh berbagai faktor seperti ketidakpuasan terhadap layanan atau produk, kurangnya keuntungan yang dirasakan, serta penawaran yang lebih baik dari pesaing [11]. Untuk itu, dibutuhkan model klasifikasi yang akurat agar bank dapat mengidentifikasi nasabah berisiko *churn* secara lebih dini. Untuk itu, dibutuhkan model klasifikasi yang akurat agar bank dapat mengidentifikasi nasabah berisiko *churn* secara lebih dini. Dalam penelitian ini, algoritma Random Forest dipilih karena mampu menangani data yang kompleks dan memberikan hasil prediksi yang stabil. Selain membangun model klasifikasi, penelitian ini juga memfokuskan evaluasi pada performa model dalam mengenali *churn* serta menganalisis fitur-fitur penting yang memengaruhi keputusan nasabah untuk berhenti menggunakan layanan.



Gambar 1. Alur Penelitian

Tahapan dalam penelitian ini dijelaskan secara sistematis melalui lima langkah utama, yang ditampilkan pada Gambar 1. Proses dimulai dari pengumpulan dan pemahaman data, dilanjutkan dengan pra-pemrosesan data untuk memastikan kualitas *input*. Setelah data siap, model Random Forest dibangun menggunakan data latih. Selanjutnya dilakukan *feature selection* untuk menyederhanakan model dan menguji pengaruh fitur terhadap performa klasifikasi. Terakhir, model dievaluasi menggunakan metrik performa untuk mengukur akurasi dan efektivitas dalam memprediksi *churn*.

1. Pengumpulan dan Pemahaman Data

Dataset yang digunakan adalah "*Churn for Bank Customers*", diperoleh dari platform Kaggle. Dataset ini terdiri dari 10.000 data nasabah dengan 13 atribut input dan 1 atribut target (*Exited*) yang menunjukkan status *churn*. Fitur-fitur mencakup usia, skor kredit, saldo, estimasi gaji, jumlah produk bank, lokasi geografis, dan status keaktifan nasabah.

2. Pra-pemrosesan Data

Fitur identifikasi seperti *RowNumber*, *CustomerId*, dan *Surname* dihapus karena tidak memiliki nilai prediktif. Fitur kategorikal *Gender* dikonversi menggunakan *label encoding*, dan *Geography* diolah dengan *one-hot encoding*. Dataset kemudian dipisah menjadi data latih dan data uji dengan rasio 70:30 menggunakan *train_test_split*. Pemisahan data menjadi dua subset ini bertujuan agar model dapat diuji pada data yang tidak pernah dilihat sebelumnya, sehingga hasil evaluasi lebih objektif. Proses encoding juga memastikan bahwa seluruh data dapat diproses oleh algoritma machine learning yang umumnya tidak dapat menangani data kategorikal secara langsung.

3. Pembangunan Model

Model dibangun menggunakan algoritma Random Forest, sebuah metode *ensemble learning* yang menggabungkan prediksi dari banyak *decision tree* yang dibuat berdasarkan subset data dan fitur yang dipilih secara acak [12]. Random Forest mampu memberikan hasil prediksi yang stabil meskipun terdapat *noise* atau ketidaksempurnaan dalam data, dan tidak rentan terhadap *overfitting* [13]. Parameter yang digunakan antara lain: *n_estimators*=200 untuk menentukan jumlah pohon yang dibangun, *min_samples_split*=5 sebagai batas minimum jumlah sampel untuk membagi node, *min_samples_leaf*=1 untuk menetapkan jumlah minimum sampel di setiap daun pohon, dan *class_weight*='balanced' untuk menangani ketidakseimbangan kelas agar model tidak bias terhadap kelas mayoritas.

4. Feature Selection

Setelah model awal dibangun, dilakukan analisis *feature importance* untuk mengidentifikasi fitur-fitur yang paling berpengaruh dalam proses klasifikasi *churn*. Random Forest secara alami menghasilkan nilai *importance* berdasarkan kontribusi tiap fitur dalam menurunkan impuritas pada *decision tree* [14]. Fitur dengan nilai *importance* lebih dari 0,03 dipilih untuk membangun model kedua, guna mengamati pengaruh penyederhanaan fitur terhadap performa klasifikasi. Fitur dengan nilai *importance* lebih dari 0,03 dipilih untuk membangun model kedua, guna mengamati pengaruh penyederhanaan fitur terhadap performa klasifikasi. Pemilihan ambang batas ini bertujuan untuk mengevaluasi apakah fitur dengan nilai *importance* yang lebih kecil benar-benar memiliki kontribusi terhadap prediksi, atau justru dapat dihilangkan tanpa menurunkan performa model secara signifikan. Dengan demikian, model kedua dapat digunakan untuk menguji apakah mempertahankan semua fitur lebih menguntungkan dibandingkan menyederhanakan model dengan hanya menggunakan fitur yang dianggap paling relevan.

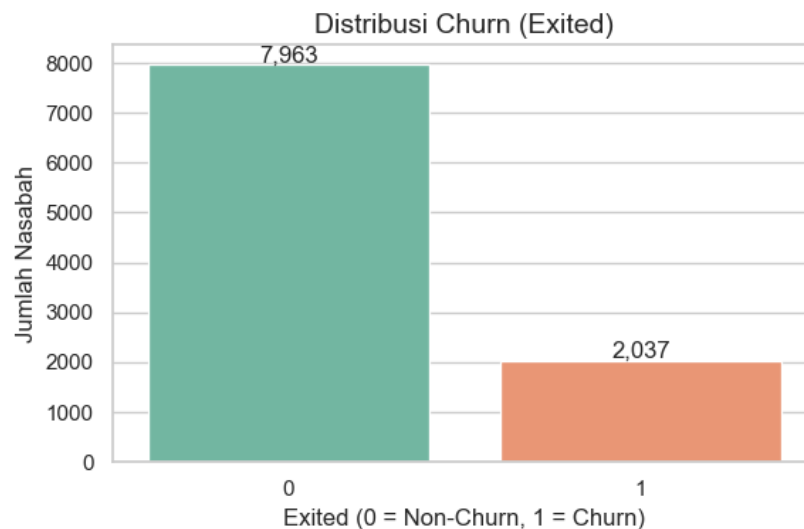
5. Evaluasi Model

Kinerja model dievaluasi menggunakan *confusion matrix* dan metrik klasifikasi seperti akurasi, *precision*, *recall*, dan *F1-score*. *Confusion Matrix* adalah metode evaluasi yang digunakan untuk membuktikan kinerja model klasifikasi dengan membandingkan hasil prediksi yang dilakukan oleh sistem dengan hasil klasifikasi yang sebenarnya [15]. *Accuracy* menunjukkan persentase prediksi yang benar dari seluruh data uji dan menggambarkan seberapa dekat hasil prediksi dengan nilai aktual [3]. *Precision* mengukur ketepatan model dalam memprediksi kelas positif, yaitu rasio antara *true positive* dan seluruh prediksi positif [16]. *Recall* menunjukkan kemampuan model dalam mengenali seluruh data positif yang sebenarnya, yaitu rasio antara *true positive* dan total aktual positif [3]. *F1-Score* merupakan rata-rata harmonis dari *precision* dan *recall*, berguna untuk menilai kinerja model pada data yang tidak seimbang [3]

3. Hasil dan Pembahasan

3.1. Hasil Pemahaman dan Karakteristik Data Awal

Dataset “Churn for Bank Customers” terdiri dari 10.000 data nasabah dengan 14 atribut seperti *Age*, *Balance*, *CreditScore*, *EstimatedSalary*, dan *NumOfProducts*. Atribut target *Exited* digunakan untuk menentukan apakah seorang nasabah *churn* (1) atau tidak *churn* (0). Hasil eksplorasi awal menunjukkan bahwa data tidak mengandung *missing value*, namun terdapat ketidakseimbangan kelas yaitu nasabah *non-churn* yang lebih banyak dibanding *churn*. Hal ini penting karena dapat memengaruhi hasil klasifikasi, dan menjadi dasar penggunaan *class balancing* pada model.



Gambar 2. Distribusi Churn

Gambar 1 menunjukkan distribusi kelas churn dalam dataset. Terlihat bahwa dari total 10.000 nasabah, sebanyak 7.963 nasabah termasuk kategori non-churn (*Exited* = 0) dan 2.037 nasabah termasuk kategori churn (*Exited* = 1). Ketidakseimbangan ini penting untuk diperhatikan karena dapat menyebabkan model cenderung bias terhadap kelas mayoritas, sehingga perlu diterapkan teknik penanganan imbalance seperti penggunaan parameter `class_weight='balanced'` pada model Random Forest. Visualisasi distribusi churn ini juga

memberikan gambaran awal mengenai tantangan utama dalam membangun model klasifikasi yang seimbang.

Atribut yang terdapat pada dataset dideskripsikan pada tabel 1 berikut.

Tabel 1: Deskripsi Atribut

No.	Atribut	Deskripsi
1.	<i>RowNumber</i>	Nomor urut data
2.	<i>CustomerId</i>	Nomor unik nasabah
3.	<i>Surname</i>	Nama belakang nasabah
4.	<i>CreditScore</i>	Skor kredit nasabah
5.	<i>Geography</i>	Negara tempat tinggal nasabah
6.	<i>Gender</i>	Jenis kelamin nasabah
7.	<i>Age</i>	Usia nasabah
8.	<i>Tenure</i>	Lama menjadi nasabah (dalam tahun)
9.	<i>Balance</i>	Saldo rekening nasabah
10.	<i>NumOfProducts</i>	Jumlah produk yang dimiliki di bank
11.	<i>HasCrCard</i>	Kepemilikan kartu kredit (1 = Ya, 0 = Tidak)
12.	<i>IsActiveMember</i>	Status keaktifan nasabah di bank (1 = Ya, 0 = Tidak)
13.	<i>EstimatedSalary</i>	Perkiraan nominal gaji nasabah
14.	<i>Exited</i>	Status <i>churn</i> (1 = <i>Churn</i> , 0 = <i>Non-Churn</i>)

3.2. Hasil Pra-pemrosesan Data

Beberapa atribut non-prediktif seperti *RowNumber*, *CustomerId*, dan *Surname* dihapus karena tidak memberikan kontribusi terhadap prediksi. Setelah penghapusan, jumlah fitur yang dipertahankan adalah 11 fitur, terdiri dari variabel numerik dan kategorikal yang telah di *encode*. Fitur *Gender* dikonversi menjadi bentuk biner (0 untuk *Female*, 1 untuk *Male*), sedangkan *Geography* diubah menjadi tiga kolom baru menggunakan teknik *one-hot encoding*, yaitu *Geography_France*, *Geography_Germany*, dan *Geography_Spain*. Hal ini menyebabkan jumlah fitur bertambah menjadi 13 fitur *input* setelah *encoding* (karena penambahan 2 kolom baru dari *Geography*, dan *Gender* tetap 1 kolom).

Dataset kemudian dipisahkan menjadi dua bagian: data latih (70%) dan data uji (30%). Dari total 10.000 data, diperoleh 7.000 baris untuk data latih dan 3.000 baris untuk data uji. Data *input* (fitur) disimpan dalam variabel X, sementara label *target* (*Exited*) disimpan dalam variabel y. Variabel X berisi data dengan 12 kolom fitur seperti *CreditScore*, *Age*, *Tenure*,

Balance, *NumOfProducts*, *EstimatedSalary*, *Geography*, *Gender*, *IsActiveMember*, dan *HasCrCard*, sedangkan *y* hanya berisi 1 kolom target yang menunjukkan status *churn* (1 = *churn*, 0 = tidak *churn*).

3.3. Hasil Penerapan Model

Model pertama dibangun menggunakan algoritma Random Forest dengan parameter yang telah dilakukan *tuning*: *n_estimators*=200, *min_samples_split*=5, *min_samples_leaf*=1, dan *class_weight*='balanced'. Model dilatih menggunakan data latih dan diuji menggunakan data uji. Hasil evaluasi model ini adalah:

Tabel 2: *Confusion Matrix* Model 1 (Tanpa Feature Selection)

	Prediksi: <i>Non-Churn</i>	Prediksi: <i>Churn</i>
Aktual: <i>Non-Churn</i>	2301 (TN)	115 (FP)
Aktual: <i>Churn</i>	283 (FN)	301 (TP)

- Akurasi: 86,7%
- *Precision churn*: 72%
- *Precision non-churn*: 89%
- *Recall churn*: 52%
- *Recall non-churn*: 95%
- *F1-score churn*: 60%
- *F1-score non-churn*: 92%

Metrik ini menunjukkan bahwa model cukup akurat dalam mengenali nasabah *churn*, meskipun *recall churn* belum terlalu tinggi, yang artinya masih ada nasabah *churn* yang luput terdeteksi.

3.4. Hasil Feature Selection

Model pertama juga menghasilkan nilai *feature importance*, yang menunjukkan kontribusi masing-masing fitur terhadap keputusan model. Fitur yang memiliki nilai *importance* di atas 0,03 adalah *Age*, *NumOfProducts*, *Balance*, *Estimated Salary*, *CreditScore*, *Tenure*, dan *IsActiveMember*. Fitur-fitur tersebut dianggap dominan karena memiliki keterkaitan langsung dengan kebiasaan dan profil keuangan nasabah. Misalnya, usia dan jumlah produk dapat mencerminkan tingkat keterlibatan nasabah, sedangkan saldo dan estimasi gaji merepresentasikan potensi nilai ekonomi dari nasabah tersebut.

Model kedua dibangun ulang hanya dengan fitur-fitur ini. Hasil evaluasi menunjukkan penurunan akurasi menjadi 85,9%, dan *recall churn* turun menjadi 46%. Ini menunjukkan bahwa beberapa fitur yang tampak “kurang penting” tetap memberi kontribusi terhadap deteksi *churn* secara keseluruhan.

Tabel 3: Confusion Matrix Model 2 (Dengan Feature Selection)

	Prediksi: <i>Non-Churn</i>	Prediksi: <i>Churn</i>
Aktual: <i>Non-Churn</i>	2308 (TN)	108 (FP)
Aktual: <i>Churn</i>	314 (FN)	270 (TP)

3.5. Hasil Evaluasi Model

Evaluasi dilakukan terhadap dua model Random Forest, yaitu model tanpa *feature selection* dan model dengan *feature selection*. Tujuan evaluasi ini adalah untuk mengetahui seberapa baik model dapat memprediksi nasabah yang berpotensi *churn*, serta untuk membandingkan efektivitas penggunaan semua fitur dengan penggunaan fitur terseleksi.

Model pertama menggunakan seluruh fitur hasil *preprocessing* dan menghasilkan akurasi sebesar 86,7%, dengan *precision* untuk kelas *churn* sebesar 72%, *recall* 52%, dan *F1-score* 60%. Nilai ini menunjukkan bahwa model cukup baik dalam mengidentifikasi nasabah yang *churn*, meskipun masih melewatkan sekitar 48% dari nasabah *churn* yang sebenarnya. Untuk kelas *non-churn*, model menunjukkan kinerja yang sangat baik dengan *precision* 89%, *recall* 95%, dan *F1-score* 92%.

Model kedua dibangun menggunakan fitur-fitur terpilih berdasarkan nilai *feature importance* > 0,03. Evaluasi menunjukkan bahwa model ini menghasilkan akurasi sebesar 85,9%, dengan *precision churn* 71%, *recall* 46%, dan *F1-score* 56%. Meski selisih akurasi tidak besar, penurunan performa pada *recall churn* cukup signifikan.

Berdasarkan hasil *feature importance*, diketahui bahwa fitur-fitur yang paling berpengaruh terhadap prediksi *churn* adalah Age, NumOfProducts, dan Balance. Nasabah dengan usia di atas 45 tahun, hanya memiliki satu produk, serta memiliki saldo tinggi merupakan kelompok yang paling berisiko *churn*. Hal ini dapat terjadi karena nasabah merasa tidak mendapatkan layanan yang sebanding dengan nilai ekonomi yang mereka miliki, atau karena tidak adanya keterikatan yang cukup dengan berbagai produk bank.

Dari temuan tersebut, bank disarankan untuk menyusun strategi retensi yang lebih terarah. Misalnya, menawarkan program loyalitas atau bundling produk untuk nasabah yang hanya menggunakan satu layanan, serta memberikan layanan prioritas atau personalisasi komunikasi kepada nasabah dengan saldo besar dan usia lebih senior. Dengan menerapkan strategi ini, bank tidak hanya dapat menekan angka *churn*, tetapi juga meningkatkan nilai hubungan jangka panjang dengan nasabah.

Untuk mempermudah perbandingan performa kedua model, berikut disajikan tabel ringkasan metrik evaluasi utama:

Tabel 4: Perbandingan Metrik Evaluasi Model

Metrik Evaluasi	Model 1 (Tanpa FS)	Model 2 (Dengan FS)
Akurasi	86.7%	85.9%
<i>Precision (Churn)</i>	72%	71%
<i>Precision (Non-Churn)</i>	89%	88%
<i>Recall (Churn)</i>	52%	46%
<i>Recall (Non-Churn)</i>	95%	96%
<i>F1-score (Churn)</i>	60%	56%
<i>F1-score (Non-Churn)</i>	92%	92%

Berdasarkan hasil evaluasi tersebut, Model 1 dipilih sebagai model terbaik, karena meskipun lebih kompleks, model ini memiliki kemampuan prediksi *churn* yang lebih baik. Selisih performa antara kedua model memang tidak terlalu besar, namun dalam praktik bisnis, kemampuan model untuk mengenali nasabah churn secara lebih akurat sangat krusial. Mendeteksi churn lebih awal memungkinkan perusahaan mengambil tindakan proaktif sebelum nasabah benar-benar meninggalkan layanan.

Selain dibandingkan dengan versi model setelah feature selection, model Random Forest dalam penelitian ini juga dibandingkan dengan penelitian terdahulu yang menggunakan metode berbeda, yaitu Naïve Bayes dan ID3. Disertakan juga perbandingan dengan penelitian terdahulu yang menggunakan dataset serta algoritma yang sama, yaitu Random Forest. Berikut disajikan tabel perbandingan performa dengan 2 penelitian terdahulu.

Tabel 5: Perbandingan Performa Penelitian Terdahulu

No	Penelitian	Algoritma	Akurasi	Keterangan
1	[3]	Naive Bayes	85.17%	Model <i>default</i>
2	[3]	ID3	79.17%	Model <i>default</i>
3	[13]	<i>Random Forest</i>	85.52%	Model <i>default</i>
4	Penelitian ini	<i>Random Forest</i>	87%	Model dengan <i>tuning parameter</i>

Berdasarkan Tabel 4.19, dapat dilihat bahwa model *Random Forest* yang dibangun dalam penelitian ini berhasil mencapai akurasi sebesar 87%, menjadikannya sebagai

model dengan performa terbaik dibandingkan dengan model-model dari penelitian terdahulu. Dua model dari penelitian [3] yaitu Naïve Bayes dan ID3 menghasilkan akurasi masing-masing sebesar 85.17% dan 79.17%. Sementara itu, model *Random Forest* dari penelitian [13] memperoleh akurasi sebesar 85.52%. Seluruh model dari kedua penelitian terdahulu tersebut menggunakan konfigurasi default tanpa proses *tuning* atau optimasi lanjutan.

4. Kesimpulan

Penelitian ini membahas penerapan algoritma Random Forest untuk membangun model klasifikasi *churn* nasabah bank dengan memanfaatkan dataset publik dari Kaggle. Dua pendekatan diterapkan, yakni model dengan seluruh fitur dan model dengan fitur terseleksi berdasarkan analisis *feature importance*. Hasil evaluasi menunjukkan bahwa model tanpa *feature selection* memiliki performa terbaik dengan akurasi sebesar 86,7% dan *F1-score* kelas *churn* sebesar 60%, sedangkan model dengan *feature selection* mengalami sedikit penurunan performa, yaitu akurasi 85,9% dan *F1-score* 56%. Hal ini menunjukkan bahwa penggunaan seluruh fitur memberikan keunggulan dalam mendeteksi *churn* secara lebih akurat, meskipun model menjadi sedikit lebih kompleks.

Analisis terhadap *feature importance* menunjukkan bahwa usia nasabah, jumlah produk yang dimiliki, dan saldo rekening merupakan faktor paling berpengaruh terhadap *churn*. Nasabah yang berusia di atas 45 tahun, hanya memiliki satu produk, serta memiliki saldo tinggi, cenderung memiliki risiko *churn* yang lebih tinggi. Temuan ini mengindikasikan bahwa *churn* tidak selalu berasal dari nasabah pasif atau berpenghasilan rendah, melainkan juga dari nasabah potensial yang merasa tidak mendapatkan pelayanan sesuai harapan.

Berdasarkan hasil tersebut, bank disarankan untuk memberikan perhatian lebih pada segmen nasabah berisiko *churn* melalui pendekatan personal, seperti penawaran produk tambahan, peningkatan layanan, atau komunikasi yang lebih proaktif. Selain itu, model klasifikasi *churn* ini dapat diintegrasikan ke dalam sistem informasi bank sebagai alat pemantauan rutin untuk mendeteksi *churn* lebih dini dan mendukung pengambilan keputusan dalam menyusun strategi retensi yang lebih efektif dan tepat sasaran.

DAFTAR PUSTAKA

- [1] Y. Yudiana, A. Y. Agustina, and N. Khofifah, "Prediksi Customer Churn Menggunakan Metode CRISP-DM Pada Industri Telekomunikasi Sebagai Implementasi Mempertahankan Pelanggan," *Indonesian Journal of Islamic Economics and Business*, vol. 8, no. 1, pp. 1–20, 2023.
- [2] D. K. Putri, F. N. Fachri, and F. T. Andini, "Pemasaran dan Riset Pemasaran Global: Konsep, Manfaat, dan Tantangan," *Jurnal Masharif Al-Syariah: Jurnal Ekonomi dan Perbankan Syariah*, vol. 9, no. 1, 2024.
- [3] M. Clementine, "Prediksi Churn Nasabah Bank Menggunakan Klasifikasi Naïve Bayes dan ID3," *Jurnal Processor*, vol. 17, no. 1, pp. 9–18, 2022.
- [4] P. P. Singh, F. I. Anik, R. Senapati, A. Sinha, N. Sakib, and E. Hossain, "Investigating customer churn in banking: a machine learning approach and visualization app for data

- science and management,” *Data Science and Management*, vol. 7, no. 1, pp. 7–16, Mar. 2024, doi: 10.1016/J.DSM.2023.09.002.
- [5] H. N. Irmanda, R. Astriratma, and S. Afrizal, “Perbandingan metode jaringan syaraf tiruan dan pohon keputusan untuk prediksi churn,” *JSI: Jurnal Sistem Informasi (E-Journal)*, vol. 11, no. 2, 2019.
- [6] T. T. Hanifa, S. Al-Faraby, and F. Informatika, “Analisis Churn Prediction pada Data Pelanggan PT. Telekomunikasi dengan Logistic Regression dan Underbagging,” *vol*, vol. 4, pp. 3210–3225, 2017.
- [7] A. Nugroho, A. Z. Fanani, and G. F. Shidik, “Evaluation of Feature Selection Using Wrapper For Numeric Dataset With Random Forest Algorithm,” in *2021 International Seminar on Application for Technology of Information and Communication (iSemantic)*, 2021, pp. 179–183. doi: 10.1109/iSemantic52711.2021.9573249.
- [8] A. Nugroho, “Analisa Splitting Criteria Pada Decision Tree dan Random Forest untuk Klasifikasi Evaluasi Kendaraan,” *JSITIK: Jurnal Sistem Informasi dan Teknologi Informasi Komputer*, vol. 1, no. 1, pp. 41–49, 2022.
- [9] R. W. Putri, A. Ristyawan, and M. N. Muzaki, “Comparison Performance of K-NN and NBC Algorithm for Classification of Heart Disease,” *JTECS : Jurnal Sistem Telekomunikasi Elektronika Sistem Kontrol Power Sistem dan Komputer*, vol. 2, no. 2, p. 143, Jul. 2022, doi: 10.32503/jtecs.v2i2.2708.
- [10] Sucipto, D. Dwi Prasetya, and T. Widiyaningtyas, “An Evaluation of the Impact of Dataset Size on Classification Performance in the Cognitive Bloom’s Taxonomy.” [Online]. Available: <https://orcid.org/0000-0003-3412-002X>
- [11] C. Mulia and A. Kurniasih, “Teknik SMOTE untuk Mengatasi Imbalance Class dalam Klasifikasi Bank Customer Churn Menggunakan Algoritma Naïve bayes dan Logistic Regression,” in *Prosiding Seminar Nasional Mahasiswa Bidang Ilmu Komputer dan Aplikasinya*, 2023, pp. 552–559.
- [12] A. Nugroho and A. Husin, “Analisis Performa Random Forest Menggunakan Normalisasi Atribut,” *SISTEMASI: Jurnal Sistem Informasi*, vol. 11, no. 1, pp. 186–196, 2022.
- [13] S. H. Hui, W. K. Khai, C. XinYing, and P. W. Wai, “Prediction of customer churn for ABC Multistate Bank using machine learning algorithms/Hui Shan Hon...[et al.],” *Malaysian Journal of Computing (MJoC)*, vol. 8, no. 2, pp. 1602–1619, 2023.
- [14] A. Orlenko and J. H. Moore, “A comparison of methods for interpreting random forest models of genetic association in the presence of non-additive interactions,” *BioData Min*, vol. 14, pp. 1–17, 2021.
- [15] K. Karsito and S. Susanti, “Klasifikasi Kelayakan Peserta Pengajuan Kredit Rumah Dengan Algoritma Naïve Bayes Di Perumahan Azzura Residencia,” *Jurnal SIGMA*, vol. 9, no. 3, pp. 43–48, 2019.

- [16] A. Nugroho and D. Harini, "Teknik Random Forest untuk Meningkatkan Akurasi Data Tidak Seimbang," *JSITIK: Jurnal Sistem Informasi dan Teknologi Informasi Komputer*, vol. 2, no. 2, pp. 128–140, 2024.